

Comparative Analysis of Self-Training XGBoost and Label Propagation for Blue Tourism Sentimen Analysis in Lampung Province

Selvya Ayu Syaputri^{*1}, Debby Alita¹

¹Departemen Of Informatics, Faculty Of Engineering And Computer Science, Universitas Teknokrat Indonesia, Bandar Lampung, Indonesia

Article details

Submitted:
July 16, 2026
Revised:
August 14, 2026
Accepted:
August 29, 2026

ABSTRACT

Blue tourism in Lampung Province is one of the tourism sectors with significant potential to support regional economic development. Tourists' perceptions of tourism destinations can be analyzed through reviews published on various digital platforms, such as Twitter/X, Instagram, and Google Maps Reviews. However, the limited availability of labeled data presents a major challenge in developing sentiment analysis models, making it necessary to employ approaches that can utilize both labeled and unlabeled data simultaneously. This research aims to compare the performance of the XGBoost-based Self-Training method and Label Propagation within a Semi-Supervised Learning approach for sentiment classification of blue tourism reviews in Lampung Province. The dataset consists of 3,947 reviews, including 500 labeled data and 3,447 unlabeled data. The research process includes text preprocessing, feature extraction using TF-IDF, partitioning labeled data into training and testing subsets with an 80:20 ratio, model training, and performance evaluation based on accuracy, precision, recall, and F1-score metrics. The trial outcomes show that XGBoost-based Self-Training reached 83% accuracy, 83.14% precision, 83% recall, and 82.98% F1-score. Conversely, Label Propagation yielded 55% accuracy, 71.31% precision, 55% recall, alongside a 50.11% F1-score. These findings demonstrate that the XGBoost-based Self-Training method outperforms Label Propagation, making it a more effective approach for sentiment classification of blue tourism reviews within a Semi-Supervised Learning framework.

Keywords : Sentiment Analysis, Semi-Supervised Learning, Blue Tourism.

ABSTRAK

Wisata biru di Provinsi Lampung merupakan sebuah sektor pariwisata yang memiliki potensi besar dalam mendukung pengembangan ekonomi daerah. Persepsi wisatawan terhadap destinasi wisata dapat dianalisis melalui ulasan yang dipublikasikan pada berbagai platform digital, seperti Twitter/X, Instagram, dan Google Maps Review. Namun, keterbatasan data berlabel menjadi tantangan dalam pengembangan model analisis sentimen sehingga diperlukan pendekatan yang mampu memanfaatkan data berlabel dan data tidak berlabel secara bersamaan. Penelitian ini bertujuan membandingkan kinerja metode Self Training berbasis XGBoost dan Label Propagation dalam pendekatan Semi-Supervised Learning untuk klasifikasi sentimen ulasan wisata biru di Provinsi Lampung. Dataset yang digunakan terdiri atas 3.947 ulasan, meliputi 500 data berlabel dan 3.447 data tidak berlabel. Tahapan penelitian meliputi preprocessing teks, ekstraksi fitur memakai Term Frequency-Inverse Document Frequency (TF-IDF), pembagian data berlabel memakai train-test split pada rasio 80:20, pelatihan model, serta evaluasi memakai metrik accuracy, precision, recall, dan F1-score. Hasil pengujian memperlihatkan bahwa metode Self Training berbasis XGBoost mendapatkan accuracy sejumlah 83%, precision sejumlah 83,14%, recall sejumlah 83%, dan F1-score sejumlah 82,98%, melainkan Label Propagation memperoleh accuracy sebesar 55%, precision sebesar 71,31%, recall sebesar 55%, dan F1-score sebesar 50,11%. Hasil penelitian menunjukkan bahwa Self Training berbasis XGBoost memiliki kinerja yang lebih baik dibandingkan Label Propagation sehingga lebih efektif diterapkan untuk klasifikasi sentimen ulasan wisata biru dalam pendekatan Semi-Supervised Learning.

Kata Kunci : Analisis Sentimen, Semi-Supervised Learning, Wisata Biru.

Corresponding author email: selayu2809@Gmail.com



This is an open access article under the [CC BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license.

1. INTRODUCTION

Tourism is one of the strategic sectors that plays an important role in promoting economic growth, increasing regional revenue, and supporting regional development in Indonesia[1]. One of the tourism sectors with significant growth potential is blue tourism, which encompasses beaches, islands, seas, and coastal areas. Lampung Province is home to various blue tourism destinations that attract both domestic and international visitors, including Pahawang Island, Gigi Hiu Beach, Kiluan Bay, Mutun Beach, and Kelagian Island[2]. The rapid development of digital platforms has encouraged tourists to share their experiences and opinions about tourism destinations through platforms such as Twitter/X, Instagram, and Google Maps Reviews. These online reviews have become valuable sources of information for evaluating destination quality, understanding tourists' perceptions of destination image, and supporting data-driven decision-making in tourism management data[3].

The large volume of tourist reviews generates massive amounts of unstructured textual data, making manual analysis inefficient and time-consuming[3]. One of the most effective approaches for processing such data is sentiment analysis, a Natural Language Processing (NLP) method designed to identify opinions as positive, negative, or neutral[4]. Sentiment analysis provides insights into tourists' perceptions of service quality, facilities, accessibility, and other factors that influence their travel experiences[3]. Most previous sentiment analysis studies have relied on Supervised Learning, necessitating a massive volume of labeled data to achieve high classification performance. However, in real-world applications, most available review data remain unlabeled, while manual annotation is costly, labor-intensive, and time-consuming. This limitation poses a significant challenge in developing sentiment classification models and makes Semi-Supervised Learning (SSL) an appropriate alternative, as it can effectively leverage both labeled and unlabeled data during the learning process [5], [6].

Among the widely adopted Semi-Supervised Learning methods are Self-Training and Label Propagation. Self-Training iteratively generates pseudo-labels for unlabeled data utilized a base classifier, whereas Label Propagation propagates label information across similar instances represented in a graph structure[7], [8]. In this study, Extreme Gradient Boosting (XGBoost) was selected as the base classifier for the Self-Training approach because its ensemble learning framework based on gradient-boosted decision trees effectively captures complex and nonlinear patterns while employing regularization to reduce overfitting, enabling the generation of more reliable pseudo-labels during the iterative self-training process and improving semi-supervised learning performance [16]. Although both methods have been widely applied in various text classification tasks[9], studies specifically comparing the performance of XGBoost-based Self-Training and Label Propagation for sentiment analysis of blue tourism reviews in Lampung Province remain limited. Furthermore, few studies have evaluated these methods using a combination of labeled and unlabeled data within the context of blue tourism in Lampung Province. This indicates the existence of a research gap that warrants further investigation[10].

Therefore, this research intends to compare the performance of XGBoost-based Self-Training and Label Propagation in classifying the sentiment of blue tourism reviews in Lampung Province using a Semi-Supervised Learning approach. The dataset consists of 3,947 tourist reviews, including 500 labeled instances and 3,447 unlabeled instances, with Term Frequency–Inverse Document Frequency (TF-IDF) employed for feature representation. Model performance is evaluated utilized accuracy, precision, recall, and F1-score. The novelty of this research lies on the comparative evaluation of XGBoost-based Self-Training and Label Propagation within a Semi-Supervised Learning framework for sentiment classification of Indonesian-language blue tourism reviews by utilizing both labeled and unlabeled data. These outcomes are anticipated to offer valuable recommendations on the best Semi-Supervised Learning approach for sentiment analysis of Indonesian text data, primarily in the tourism sector.

2. RESEARCH METHOD

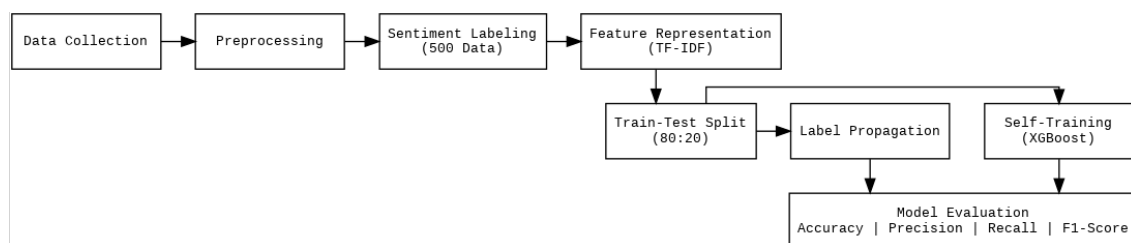


Fig. 1. Research Method

This study employed an experimental approach to compare the performance of two Semi-Supervised Learning methods, namely XGBoost-based Self-Training and Label Propagation, in classifying the sentiment of blue tourism reviews in Lampung Province. A quantitative approach was applied by evaluating the performance of both methods utilized Accuracy, Precision, Recall, and F1-Score as evaluation metrics. The research stages consisted of data collection, data preprocessing, data labeling, feature extraction using Term Frequency–Inverse Document Frequency (TF-IDF), model training, and classification performance evaluation.

2.1 Data Collection

The dataset was collected from three digital platforms, namely Instagram, Google Maps, and X (formerly Twitter), using the Apify Console web scraping service. Data collection was performed using keywords and hashtags related to blue tourism destinations in Lampung Province. The collected data consisted of review texts along with the available supporting information from each platform. All collected data were merged into a single CSV dataset for preprocessing. A total of 3,947 reviews were successfully collected.

Table 1. Dataset Distribution

Platform	Jumlah data
Instagram	273
Google Maps	2.733
Twitter/X	941

2.2 Data Preprocessing

Data preprocessing represents the starting step that intends to clean raw data for analysis. At this stage, the collected data were processed through a series of cleaning and transformation techniques to obtain a consistent structure suitable for sentiment analysis. Preprocessing was performed to reduce noise, eliminate inconsistencies, and improve text representation, thereby enhancing the effectiveness of the data mining process. The preprocessing steps included case folding, text cleaning, tokenization, stopword removal, and stemming.

- 1) *Case Folding* : Case folding converts all characters in each review into lowercase letters using the `.lower()` function. This process standardizes word representation so that differences in letter capitalization do not affect feature extraction or text analysis.
- 2) *Cleaning Text* : The text cleaning stage removes characters that do not contribute to sentiment analysis. URLs, user mentions, hashtags (#), numbers, and punctuation marks were removed using regular expressions through the `re.sub()` function and the `string.punctuation` library. In addition, leading and trailing whitespace was removed using the `.strip()` function.
- 3) *Tokenizing* : Tokenization separates each review into individual words (tokens) using the `word_tokenize()` function provided by the Natural Language Toolkit (NLTK) library. This process facilitates subsequent preprocessing steps, particularly stopword removal and stemming.
- 4) *Stopword Removal* : Stopword removal eliminates frequently occurring words that provide little semantic information, such as conjunctions, prepositions, and other common words. This process utilizes the Indonesian stopword list provided by the NLTK library and is adjusted according to the research requirements.
- 5) *Stemming* : Stemming converts inflected words into their root forms using the Nazief–Adriani stemming algorithm implemented through the Sastrawi library. This process reduces variations of words with similar meanings, resulting in a more consistent text representation.

2.3 Sentimen Labeling

After preprocessing, 500 reviews were manually labeled to serve as the initial labeled dataset. Each review was allocated to one of three sentiment categories: positive, neutral, or negative, based on the meaning and opinion expressed in the text. The remaining 3,447 reviews were treated as unlabeled data and utilized during the Semi-Supervised Learning process. The labeled data were utilized to build the initial classification model, whereas the unlabeled data were exploited by the XGBoost-based Self-Training and Label Propagation methods to generate pseudo-labels and expand the training dataset automatically.

2.4 TF-IDF Feature Representation and Train-Test Split (80:20)

Following preprocessing and sentiment labeling, the labeled dataset was split into training and testing portions through an 80:20 Train-Test Split with stratified sampling implemented in the Scikit-learn library to preserve the class distribution. Consequently, 400 reviews were allocated to the training set and 100 reviews to the testing set, while the 3,447 unlabeled reviews remained available for the Semi-Supervised Learning process.

Data were converted into numerical feature vectors utilizing the Term Frequency-Inverse Document Frequency (TF-IDF) representation implemented through `TfidfVectorizer` with `max_features = 5,000`. The TF-IDF vocabulary was generated exclusively from the training data, whereas both the testing and unlabeled datasets were transformed using the same feature space to maintain consistency during model training, testing, and pseudo-label generation [11], [12].

2.5 Xgboost-Based Self Training

Self-Training is a Semi-Supervised Learning method that employs both labeled and unlabeled data during the learning process. The method begins by training an initial classifier using the labeled dataset. The trained classifier then predicts sentiment labels for the unlabeled instances together with their confidence scores. Predictions with confidence scores exceeding a predefined threshold are considered reliable and are assigned as

pseudo-labels. These pseudo-labeled instances are subsequently incorporated into the labeled training dataset, allowing the classifier to be retrained iteratively using the expanded dataset until no additional high-confidence predictions are obtained. In this study, XGBoost was employed as the base classifier to construct the initial model utilizing the labeled training data.

The XGBoost classifier was configured with $n_estimators = 100$, $random_state = 42$, and $eval_metric = "mlogloss"$. It was implemented as the base estimator of SelfTrainingClassifier with a confidence threshold of 0.8. Finally, the trained model was assessed utilizing the testing dataset.

2.6 Label Propagation

Label Propagation is a graph-based Semi-Supervised Learning technique that transmits label information from labeled nodes to unlabeled nodes based on their similarity relationships. The algorithm constructs a similarity graph in which each data instance is represented as a node, while edges indicate the similarity between instances. During the learning process, label information is iteratively propagated from labeled nodes to neighboring unlabeled nodes until the label assignments converge. Consequently, unlabeled instances receive labels according to the influence of nearby labeled instances with similar feature representations. In this study, the algorithm was implemented using the LabelPropagation class from the Scikit-learn library. The method utilized both labeled and unlabeled TF-IDF feature vectors. The Label Propagation model employed the K-Nearest Neighbor (KNN) kernel with $n_neighbors = 7$ and $max_iter = 1000$. The $n_neighbors$ parameter specifies the number of nearest neighbors involved in label propagation, whereas max_iter determines the maximum number of iterations before convergence. The resulting predictions were evaluated utilizing Accuracy, Precision, Recall, and F1-score.

2.7 Model Evaluasi

Model performance was evaluated by comparing the predicted sentiment labels with the ground-truth labels of the testing dataset. Four performance metrics were utilized: Accuracy, Precision, Recall, and F1-score. Accuracy quantifies the ratio of rightly categorized samples across total test instances. Precision evaluates the ratio of rightly forecasted instances for every designated category. Recall gauges the model's capacity to rightly detect every pertinent instance of any category. The F1-score denotes the harmonic average of Precision and Recall, offering an equitable evaluation of predictive performance. Because the sentiment classification problem involves three categories (positive, neutral, and negative), weighted average metrics were chosen for Precision, Recall, and F1-score to address the class imbalance issue.

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (1)$$

$$Precision = \frac{TP}{TP+FP} \quad (2)$$

$$Recall = \frac{TP}{TP+FN} \quad (3)$$

$$F1-Score = \frac{2 \times Precision \times Recall}{Precision + Recall} \quad (4)$$

3. RESULTS AND DISCUSSION

3.1. Implementasi of the Self-Training (XGBoost) Method

In this research, the Self-Training method was implemented using the XGBoost algorithm as the base classifier. The model was trained using 400 labeled data as the training set and 3,447 unlabeled instances within this semi-supervised learning processes. Subsequently, the trained model was used to classify 100 testing data that had been separated during the Train-Test Split stage. The experimental results of the Self-Training (XGBoost) method achieved an Accuracy of 0.8300, Precision of 0.8314, Recall of 0.8300, and F1-Score of 0.8298.

Table 2. Evaluation Results of the Self-Training (XGBoost) Method

Metric	Value
Accuracy	0,8300
Precision	0,8314
Recall	0,8300
F1-Score	0,8298

In addition to the evaluation metrics, the classification results of the Self-Training (XGBoost) method are also presented using a Confusion Matrix to illustrate the prediction distribution for each sentiment class.

Table 3. Confusion Matrix Results

Actual \ Predicted	Negative	Neutral	Positive
Negative	29	3	1
Neutral	3	28	3
Positive	3	4	26

3.2. Implementation of the Label Propagation Method

In this study, the Label Propagation method was implemented using the K-Nearest Neighbor (KNN) kernel with $n_neighbors = 7$ and $max_iter = 1000$. The model was trained using 400 labeled data and 3,447 unlabeled data represented by TF-IDF features. Subsequently, the model was used to classify 100 testing data separated during the Train-Test Split stage. The experimental results of the Label Propagation method achieved an Accuracy of 0.5500, Precision of 0.7131, Recall of 0.5500, and F1-Score of 0.5011.

Table 4. Evaluation Results of the Label Propagation Method

Metric	Value
Accuracy	0,5500
Precision	0,7131
Recall	0,5500
F1-Score	0,5011

In addition to the evaluation metrics, the classification results of the Label Propagation method are also presented using a Confusion Matrix to illustrate the prediction distribution for each sentiment class.

Table 5. Confusion Matrix Results of the Label Propagation Method

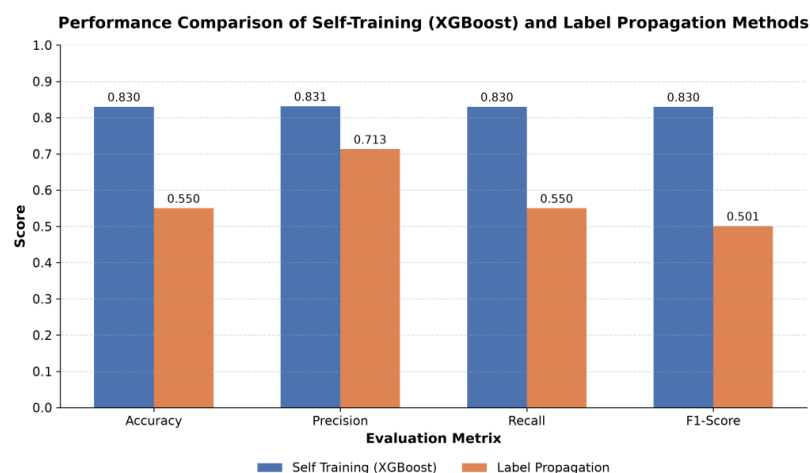
Actual \ Predicted	Negative	Neutral	Positive
Negative	29	0	4
Neutral	23	5	6
Positive	12	0	21

3.3. Evaluation and Performance Comparison of the Methods

The evaluation was conducted to compare the performance of the Self-Training (XGBoost) and Label Propagation methods in classifying the sentiment of blue tourism reviews. The comparison was referring on four evaluation metrics: Accuracy, Precision, Recall, and F1-Score.

Table 6. Evaluation Results of Both Methods

Method	Accuracy	Precision	Recall	F1-Score
Self Traing (XGBoost)	0,8300	0,8314	0,8300	0,8298
Label Propagation	0,5500	0,7131	0,5500	0,5011

**Fig. 2.** Performance Comparison of the Self-Training (XGBoost) and Label Propagation Methods

Based on the results, the Self-Training (XGBoost) method demonstrated better performance than Label Propagation across all evaluation metrics. Self-Training achieved an Accuracy of 0.83, whereas Label Propagation

achieved 0.55. Furthermore, the Precision, Recall, and F1-Score obtained by Self-Training were also higher than those of Label Propagation.

The superior performance of the Self-Training method indicates that the use of XGBoost as the base classifier is able to utilize the combination of labeled and unlabeled data more effectively. During the Self-Training process, the model gradually assigns labels to unlabeled data with high confidence, and these newly labeled data are subsequently incorporated into the training set to update the model, thereby improving its classification capability.

In contrast, the Label Propagation method produced lower performance on the dataset used in this study. This indicates that the label propagation process based on similarity among data instances was unable to construct an optimal class representation. This condition may be influenced by the characteristics of the review data, which contain a high diversity of vocabulary, making the relationships among data instances difficult to represent solely through feature similarity.

Referring on the evaluation outcome, it can be concluded that the Self-Training (XGBoost) method is the most effective approach for sentiment classification of blue tourism reviews in this study, achieving an Accuracy of 83%, Precision of 83.14%, Recall of 83%, and an F1-Score of 82.98%.

3.4. Application of the Model to Unlabeled Data

After the training and evaluation processes of both methods were completed, the models were applied to the 3,447 review data that had not yet been assigned sentiment labels. This stage aimed to automatically generate sentiment labels using the Self-Training (XGBoost) and Label Propagation methods so that all data could be utilized for subsequent analysis.

Table 7. Sentimen Distribution Generated by the Self-Training (XGBoost) Method

Sentiment	Number of Data
Positive	1.094
Neutral	1.859
Negative	494
Total	3.447

Based on these results, the Self-Training (XGBoost) method produced predictions dominated by the neutral sentiment class with 1,859 data (53.93%), next came the positive sentiment class with 1,094 cases (31.74%), while the negative sentiment class accounted for the smallest proportion with 494 data (14.33%). This distribution indicates that most unlabeled blue tourism reviews tend to express neutral to positive sentiments. Subsequently, the Label Propagation method was also applied to the same dataset.

Table 8. Sentimen Distribution Generated by the Label Propagation Method

Sentiment	Number of Data
Positive	634
Neutral	263
Negative	2550
Total	3.447

Unlike Self-Training, the Label Propagation method generated a distribution dominated by the negative sentiment class with 2,550 data (73.98%), while the positive sentiment class consisted of 634 data (18.39%) and the neutral sentiment class contained only 263 data (7.63%). This difference indicates that Label Propagation tends to assign predictions to a single dominant class, namely the negative class. The difference in labeling results between the two methods is consistent with the previous evaluation results. The Self-Training (XGBoost) method achieved an Accuracy of 83%, whereas Label Propagation achieved only 55%. Therefore, the automatic labeling results produced by the Self-Training (XGBoost) method are considered more representative for subsequent sentiment analysis of unlabeled data, as they provide better classification performance and a more balanced prediction distribution than Label Propagation[13], [14], [15].

4. KESIMPULAN

This study successfully implemented Semi-Supervised Learning using the Self-Training (XGBoost) and Label Propagation methods for sentiment classification of blue tourism reviews in Lampung Province. Both methods utilized 400 labeled data and 3,447 unlabeled data during the training process, while the evaluation was conducted using 100 testing data. Experimental results demonstrated that the XGBoost-based Self-Training method consistently outperformed Label Propagation across all evaluation metrics. Self-Training also produced a more balanced sentiment distribution on the unlabeled dataset, whereas Label Propagation tended to assign a larger proportion of instances to the negative sentiment class.

These findings suggest that the iterative pseudo-labeling mechanism of Self-Training enables XGBoost to

utilize information from both labeled and unlabeled data more effectively, thereby improving its ability to distinguish sentiment classes. In contrast, Label Propagation relies on similarity-based label propagation, which may result in less discriminative classifications when the similarity structure of the data is not sufficiently representative. Therefore, the XGBoost-based Self-Training method can be considered a more suitable Semi-Supervised Learning approach for sentiment classification of blue tourism reviews in Lampung Province. Its ability to effectively exploit unlabeled data while maintaining higher classification performance demonstrates its potential for supporting large-scale sentiment analysis, particularly in situations where manually labeled data are limited.

DAFTAR PUSTAKA

- [1] H. Fadilla, "Pengembangan Sektor Pariwisata untuk meningkatkan Pendapatan Daerah Di Indonesia," *J. Business, Econ. Financ.*, vol. 2, 2024, doi: 10.31080/BENEFIT.
- [2] S. Rahmadila and D. Alita, "Perbandingan Metode Naive Bayes Classifier dan Support Vector Machine Pada Analisis Sentimen Wisata Biru Berdasarkan Ulasan Twitter, Instagram, dan Google Maps Review," *Technol. Sci.*, vol. 7, no. 3, 2025, doi: 10.47065/bits.v7i3.8791.
- [3] X. Guo and J. A. Pesonen, "The Role of Online Travel Reviews in Evolving Tourists' Perceived Destination Image," *Scand. J. Hosp. Tour.*, vol. 22, no. 4–5, pp. 372–392, 2022, doi: 10.1080/15022250.2022.2112414.
- [4] B. Liu, J. Zhao, K. Liu, and L. Xu, *c. Press*, 2015. doi: 10.1162/COLI.
- [5] J. E. van Engelen and H. H. Hoos, "A Survey on Semi-supervised Learning," *Mach. Learn.*, vol. 109, no. 2, pp. 373–440, 2020, doi: 10.1007/s10994-019-05855-6.
- [6] I. G. S. M. U. F. S. V. H. S. H. A. F. Dyasa, *Machine Learning*. Media Nusa Creative (MNC Publishing), 2026.
- [7] M. W. Dunham, A. Malcolm, and J. K. Welford, "Improved Well-log Classification Using Semisupervised Label Propagation and Self-training, with Comparisons to Popular Supervised Algorithms," *Geophysics*, vol. 85, pp. O11–O115, 2020, doi: 10.1190/GEO2019-0238.1.
- [8] H. Liu, S. K. Rallabandi, Y. Wu, P. P. Dakle, and P. Raghavan, "Self-training Strategies for Sentiment Analysis: An Empirical Study," in *Findings of the Association for Computational Linguistics*, 2024.
- [9] T. Yang, L. Hu, C. Shi, H. Ji, X. Li, and L. Nie, "HGAT: Heterogeneous Graph Attention Networks for Semi-supervised Short Text Classification," *ACM Trans. Inf. Syst.*, vol. 39, no. 3, 2021, doi: 10.1145/3450352.
- [10] C. Agustina, P. Purwanto, and F. Farikhin, "Enhancing Sentiment Analysis Accuracy in Borobudur Temple Visitor Reviews through Semi-Supervised Learning and SMOTE Upsampling," *J. Adv. Inf. Technol.*, vol. 15, no. 4, pp. 492–499, 2024, doi: 10.12720/jait.15.4.492-499.
- [11] R. Kusumaningrum, A. A. Herlambang, W. Afifah, A. Wibowo, Sutikno, and P. S. Sasongko, "Semi-Supervised Learning vs. Few-Shot Learning: Which is Better for Sentiment Analysis on Hotel Reviews Towards a Small Labeled Training Data?," *Int. J. Adv. Comput. Sci. Appl.*, vol. 16, no. 11, pp. 671–677, 2025, doi: 10.14569/IJACSA.2025.0161168.
- [12] D. Guidotti, L. Pandolfo, and L. Pulina, "Discovering sentiment insights: streamlining tourism review analysis with Large Language Models," *Inf. Technol. Tour.*, vol. 27, pp. 227–261, 2025, doi: 10.1007/s40558-024-00309-9.
- [13] C. Xu, M. Wang, and S. Zhu, "Tourism sentiment quadruple extraction via new neural ordinary differential equation networks with multitask learning and implicit sentiment enhancement," *Expert Syst. Appl.*, vol. 270, p. 126417, 2025, doi: 10.1016/j.eswa.2025.126417.
- [14] M. Hajiabadi, H. Vahdat-Nejad, and H. Hajiabadi, "Enhancing tourism sentiment analysis with deep learning: a comprehensive study on social media data," *Curr. Issues Tour.*, 2025, doi: 10.1080/13683500.2025.2583465.
- [15] C. R. Utami and I. Santiko, "Comparison of Support Vector Machine and XGBoost Algorithms in Sentiment Analysis of Visitor Reviews of Baturaden Tourism Forest," *J. Multimed. Trend Technol.*, vol. 4, no. 2, 2025.
- [16] H. Wang, Q. Sun, J. Wu, X. Zhang, W. Liu, T. Peng, and R. Tang, "Self-training-based approach with improved XGBoost for aluminum alloy casting quality prediction," *Robot. Comput.-Integr. Manuf.*, vol. 92, Art. no. 102890, 2025, doi: 10.1016/j.rcim.2024.102890.